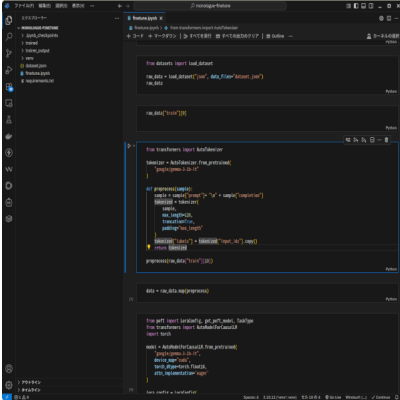




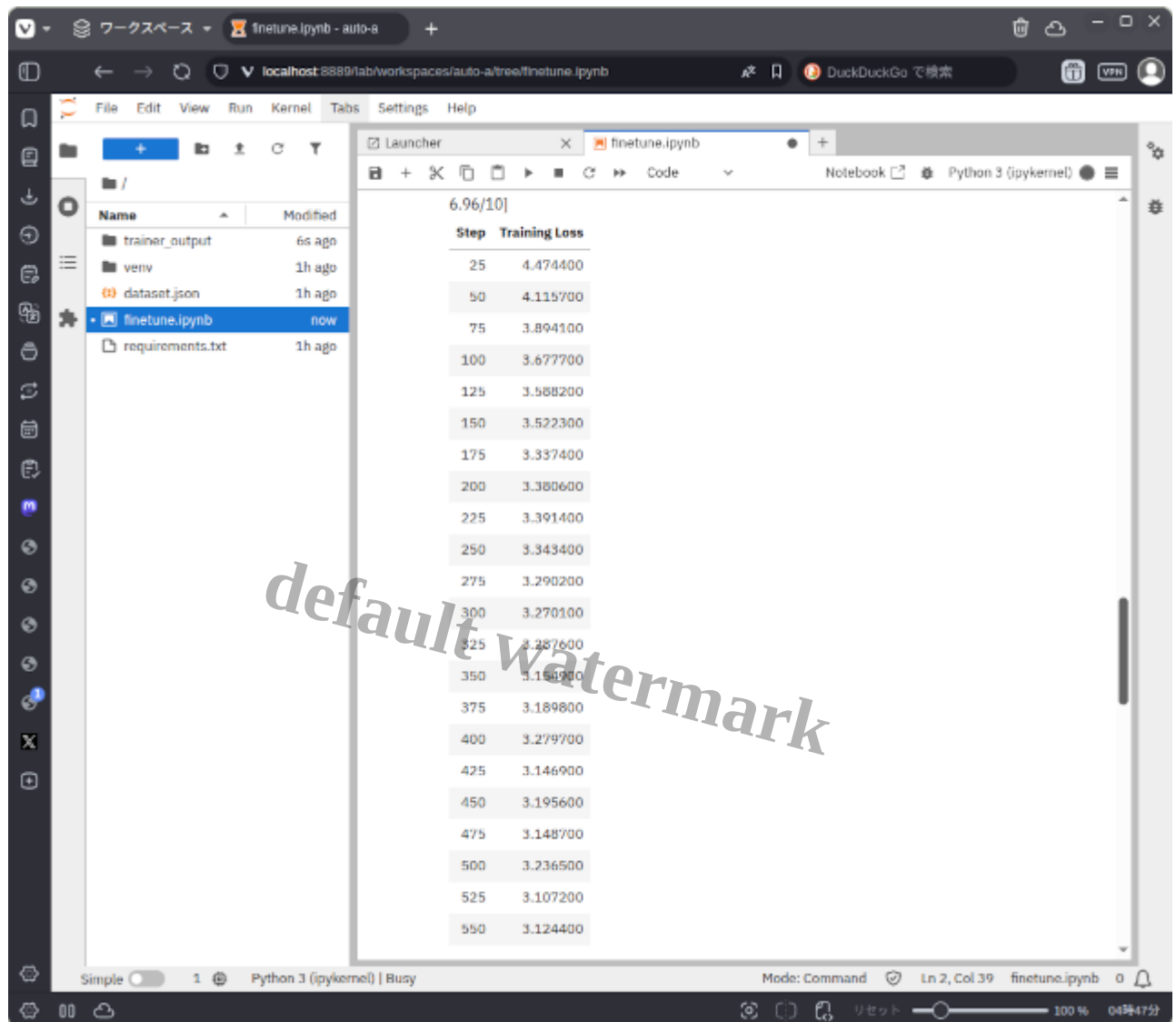
ç?¥è?è??ç©•ã?»±æ??ã•®ç ?ç©¶ã??ã•ã•®i¼? LLMã??ã?jã?ã?³ã?
ã?¥ã?¼ã??ã?³ã?°ã•®è©|ã•¿

Description

å ?æ?¥ã?•ã•®ã??ã?°ã??ã?¬ã?ã?¼ã?³ã?°ã?ã?æ??ç?¿ã??æ©?æ¢°å¿çç?ç?¬ã•
®ã??ã?¼ã?¿ã?»ã??ã??ã•«å±?æ•?ã•?ã??ã??ã?ã?°ã?©ã?ã??æ?ã•ã•¼ã•?ã•?ã??ã?»å??ã¬ã•ã•
®ç¶?ã•ã•§ã?•ã•ã•®ã??ã?¼ã?¿ã?»ã??ã??ã??ã½¿ã•£ã•?æ?¢ã?ã•®å-
¿ç¿æ¿?è?è°ã?ã?¢ã??ã?«i¼?Pretreined Language Modeli¼?ã??Loraã?•ã?¥ã?¼ã??ã?³ã?°ã•
?ã??ã??ã?ã?ã?©ã?ã??æ?ã•?ã?ã?¿ã?¿ã•¼ã•?ã•?ã??i¼?ã¿ã•®ã?¹ã¬¬ã?ã?¼ã?³ã?ã?§ã??ã??i¼?

ã½¿ç?¬ã?ã?ã?ã?³ã??ã?¥ã?¼ã?¿ã?¼ã¬Core i5 13500 CPUã¬NVIDIA RTX A4000 (16GB
VRAM)ã??æ•?¼?ã•?ã?ã?ã?ã•®ã??ã¬¼è±jã¬ã?ã?è?è°ã?ã?¢ã??ã?«ã¬Googleã®Gemma 3 1Bã•
®ã?ã?³ã?¹ã??ã?©ã¬¬ã?ã?§ã?³ã?ã?¥ã?¼ã??ã?³ã?°æ¿ã•®ã??ã•®ã•§ã?ã??

çµæ??ã¬å±±æ??ã??å¿ç¿ã?ã?¼ã?æ¿ã??ã?ã®ç?ç?ç?ã•ã??ã•®ã¬«ã¬ã?ã??ã•¼ã•?ã??ã?§ã?ã•
?ã??



ã•?ã•®ç`?â!ã•®â?|ç•?ã•§ã•?ã•`ã•?ã•³ã•³ã•?ã•¥ã•¼ã•¿çã•¼ã•®â?|ç•?é??â!ã•«ã•§ã•ã•²ã•æ°?ã•-ã•ã•
?ã•?ã•®ã•®ã•VRAMã•-ã•?ã•£ã•æ-²ã•?ã•?ã•§ã•?ã•?ã•?ã•®ã•³ã•³ã•?ã•¥ã•¼ã•¿çã•¼ã•®æ§?æ?ã•
§ã•-1Bã•®ã•?£ã•?ã•?ã•?ã•?ã•¼ã•?-ã•³é•.128ã•ã•?ã•?ã•?é?•ç??ã•§ã•?14GBã•»ã•©ã•
®VRAMã•?ã½çã•£ã•!ã•?ã¾ã•?ã•?ã•?

ã?ã?ã??ã?¼ã??ã?©ã?|ã?¼ã?¿ã?¼ã?_ã½?ã°|ã??ã¤?æ?´ã?ã?|è©|ã?ã?¾ã?ã?ã?ã?ã?è?_ã?çµ•
æ??ã?_ã¾?ã??ã??ã¾ã?ã??ã?šã?ã?ã??

ã??ã?¼ã?¿ã?»ã??ã??ã·®è³ªã·æ?ªã·?¼?ã??ã?ã°æ???¿¿¿ã·ã·®ã·¾ã·¾ã¯æ??ã¾ã·ã·ã·ã·?ï¼?ã·
“ã·ã·ã·ã·ã·ã·ã·ã·ã·ã·ã·ã·¾ã·ã·ã·

å•?è??ã¼ã•šã•«ã?•ã??ã?-ã?¼ã??ã?³ã?°ã?³ã?¼ã??ã-æ-ja®ã??ã?•ã•ã??ã®ã•šã•?ã??

/* ã??ã?¼ã?¿ã?»ã??ã??ã??è^aã•¿:è^{¾¼}ã?? */

```
from datasets import load_dataset
```

```
raw_data = load_dataset("json", data_files="dataset.json")
```

```
/* ä??ä?¼ä?¬ä??ä?¼ä?¶ä?¼ä?•@æ°?å?? */  
  
from transformers import AutoTokenizer  
  
tokenizer = AutoTokenizer.from_pretrained(  
    "google/gemma-3-1b-it"  
)  
  
/* ä??ä?¼ä?¿ä?»ä??ä??ä?•@å?•å?|ç?• */  
  
def preprocess(sample):  
    sample = sample["prompt"]+ "\n" + sample["completion"]  
    tokenized = tokenizer(  
        sample,  
        max_length=128,  
        truncation=True,  
        padding="max_length"  
    )  
    tokenized["labels"] = tokenized["input_ids"].copy()  
    return tokenized  
  
data = raw_data.map(preprocess)  
  
/* ä?•ä?¥ä?¼ä??ä?³ä?°å¬¼è±¼ä?•@ä?¢ä??ä?«ä?•@æ°?å?? */  
  
from peft import LoraConfig, get_peft_model, TaskType  
from transformers import AutoModelForCausalLM  
import torch  
  
model = AutoModelForCausalLM.from_pretrained(  
    "google/gemma-3-1b-it",  
    device_map="cuda",  
    torch_dtype=torch.float16,  
    attn_implementation='eager'  
)  
  
/* LORAä?•@è¨å@? */  
lora_config = LoraConfig(  
    task_type=TaskType.CAUSAL_LM,  
    target_modules=["q_proj", "k_proj", "v_proj"]  
)  
  
model = get_peft_model(model, lora_config)  
  
/* ä??ä?-ä?¼ä??ä?¼ä?•@æ°?å?? */  
  
from transformers import TrainingArguments, Trainer  
  
train_args = TrainingArguments(  
    num_train_epochs=10,  
    learning_rate=3e-4,  
    logging_steps=25,  
    fp16=True  
)
```

```
trainer = Trainer(
    args=train_args,
    model=model,
    train_dataset=data["train"]
)

/* ã??ã?-ã?¼ã??ã?³ã?°å®?èì? */
trainer.train()

/* ã??ã?-ã?¼ã??ã?³ã?°æ_ ?ã•®ã?¢ã??ã?«ã•¨ã??ã?¼ã?¬ã??ã?ðã?¶ã?¼ã•®ä¿•å? */

trainer.save_model("./trained")
tokenizer.save_pretrained("./trained")

!ç¿?ã?ã?¹ã¬1.0ä»¥ä_ã?ç?®æ¨?ã•§ã?ã?ã?•3.0ã¾ã•§ã?ã?æ_ã??ã¾ã?ã??ã•§ã?ã?ã??
ä»¥ä_ã?®ã³ã?¼ã??ã?ã??ã?;ã?ðã³ã?ã?¥ã¼ã??ã³ã?°å¾¾ã?®ã?¢ã??ã?«ã??ã½¿ã•£ã•
?æ?“è«ã??èì?ã?ã??ã•®ã•§ã?ã??
```

Category

1. æ?¥ã? ã•®æ'»å??

Tags

1. AI
2. Linux
3. LLM
4. Lora
5. ã??ã?jã?□ã?³ã?•ã?¥ã?¼ã??ã?³ã?°
6. ä,?é??å??å,?
7. å□Sè!•æ"jè"èª?ã?¢ã?ã?«
8. å±±æ¢"ç??

Date Created

2025å¹¸æ??3æ?¥

Author

kazu0-tsubaki